

Research Paper

Bayesian prediction for the exponential distribution under the exponential squared error loss function

SEYED MOHAMMAD TAGHI KAMEL MIRMOSTAFAEI*, REZA GHASABANI
DEPARTMENT OF STATISTICS, UNIVERSITY OF MAZANDARAN, P.O. BOX 47416-1467,
BABOLSAR, IRAN

Received: May 04, 2025/ Revised: November 10, 2025/ Accepted: January 06, 2026

Abstract: A recently proposed asymmetric loss function, termed the exponential squared error loss function, has been applied to estimation problems. However, prediction methodologies based on this loss function remain unexplored in the literature. In this paper, we focus on Bayesian two-sample prediction for the exponential distribution under this new loss function. We consider situations where the informative sample size may be either fixed or random. Since the Bayesian predictors do not have closed forms, we employ a Monte Carlo method to approximate them. A simulation study is conducted to evaluate the performance of the predictors. The simulation results indicate that predictors with fixed sample sizes can outperform those with truncated Poisson and truncated geometrically distributed sample sizes. A real data example is also presented for illustration. The paper concludes with a discussion.

Keywords: Bayesian point predictor; Exponential squared error loss function; Monte Carlo method; Random sample size; Simulation.

Mathematics Subject Classification (2010): 62F15, 65C05.

1 Introduction

The exponential distribution is widely recognized as one of the most fundamental lifetime distributions in reliability theory. Let X follow an exponential distribution with parameter θ . Then, the probability density function (pdf) of X is given by

$$f(x; \theta) = \theta \exp(-\theta x), \quad x > 0, \theta > 0, \quad (1)$$

*Corresponding author: m.mirmostafaei@umz.ac.ir

and the cumulative distribution function of X is given by

$$F(x; \theta) = 1 - \exp(-\theta x), \quad x > 0, \theta > 0,$$

and we may write $X \sim \text{Exp}(\theta)$. Numerous studies have established their methodologies using the one-parameter exponential distribution as the base distribution, see for example, Asgharzadeh and Valiollahi (2012), Ahmadi et al. (2016), Kumar et al. (2020), and Fallah and Asgharzadeh (2021).

Prediction of future order statistics has been extensively studied using both classical and Bayesian approaches. For instance, Lawless (1971), and Lawless (1977) presented techniques for predicting unobserved order statistics from the exponential distribution based on the observed order statistics. Kaminsky and Nelson (1975) worked on the best linear unbiased predictors of order statistics in location and scale families. Lingappaiah (1978) applied the Bayesian framework to the prediction problem for the exponential population. Lingappaiah (1979) focused on the prediction of future order statistics in exponential and gamma distributions using the Bayesian methodology. MirMostafae and Ahmadi (2011) discussed the point prediction of future statistics from the exponential distribution based on observed record values. Asgharzadeh and Valiollahi (2012) considered the prediction of times to failure of censored units in hybrid censored samples from an exponential distribution. The above-mentioned investigations have assumed that the sample sizes are fixed. However, many research scenarios cannot guarantee fixed sample sizes due to participant dropout, measurement errors, or other factors, see also Srivastava (1973). Numerous studies have adopted the framework where the sample size is treated as a discrete random variable. Early theoretical developments regarding the distributions of order statistics under this framework were established by Raghunandan and Patil (1972), Buhrman (1973), Consul (1984), Gupta and Gupta (1984) and Rohatgi (1987). Many authors investigated the prediction problem when the sample size is random. For example, Lingappaiah (1990) studied the Bayesian prediction of the order statistics coming from an exponential distribution when an outlier is present in the sample drawn and the sample size is a random variable. Ashour and El-Wekeel (1993) investigated Bayesian prediction of a future order statistic in the generalized Burr distribution under a random future sample size. Nigm and Abd Al-Wahab (1996) developed Bayesian prediction bounds for the Burr distribution with a discrete random sample size. Soliman (2000) examined the one-sample prediction problem under a random sample size, when the underlying model is the Pareto distribution. AL-Hussaini and Al-Awadhi (2010) obtained Bayesian two-sample predictors for generalized order statistics, considering both fixed and random future sample sizes. Basiri and Ahmadi (2015) worked on the prediction of future generalized order statistics when the future sample size is taken as a random variable. Ahmadi et al. (2016) optimized sample sizes using a cost function in two-sample Bayesian prediction, taking into account both fixed and random information sample sizes. Shafay et al. (2017) explored the Bayesian prediction of order statistics given k -record values from a Pareto distribution, examining fixed and random sample sizes. Barakat et al. (2018) developed pivotal quantities for constructing prediction intervals for future lifetimes that follow a two-parameter exponential distribution, based on a random number of generalized order statistics. Basiri and Pakzad (2018) analyzed the problem of the prediction of order statistics based on record values, considering fixed and random future sam-

ple sizes. Barakat et al. (2021) proposed an efficient point predictor for future order statistics when the sample size is treated as a random variable.

The problem of prediction for the exponential distribution under well-known loss functions, such as the squared loss function, has been addressed by many authors, see for example, Ahmadi et al. (2016). Recently, a new asymmetric loss function has been proposed by Kumar et al. (2020), which is called the exponential squared error loss (ESEL) function. This loss function is defined as

$$L = (e^{-\delta} - e^{-\theta})^2. \quad (2)$$

The estimator of θ under the ESEL function is given by (Kumar et al., 2020)

$$\hat{\theta} = -\log E(\exp(-\theta)|data).$$

Kumar et al. (2020) employed the ESEL function to obtain estimators for the parameter and for the reliability function of the exponential distribution. Later, Kumar et al. (2023) implemented the ESEL function to derive an estimator for the parameter of the Minimum Guarantee exponential distribution. However, the prediction problem under this new loss function has not been addressed yet.

This paper aims to find predictors for a future ordered observation coming from the exponential distribution based on a random sample with fixed and random sizes under the ESEL loss function. Since the predictors do not possess closed forms, we propose a Monte Carlo method to approximate these predictors. In what follows, first, we present the process of obtaining the predictor of a future unobserved order statistic from $Exp(\theta)$ when the sample size is fixed in Section 2. Section 3 is devoted to the results when the sample size is a random variable and its distribution belongs to the power series family of distributions. An extensive simulation is conducted and presented in Section 4. A real data example is analyzed in Section 5. The paper ends with a discussion.

2 Main results when $N = n$ is fixed

Let $Y_{s:m}$ represent the s -th future order statistic from a sample of size m , where the sample consists of independent and identically distributed (i.i.d.) continuous random variables following a one-parameter exponential distribution with pdf (1). Then, the pdf of $Y_{s:m}$ is given by (Arnold et al., 2008)

$$f_{Y_{s:m}}(y; \theta) = \frac{\theta e^{-(m-s+1)\theta y}}{B(s, m-s+1)} (1 - e^{-\theta y})^{s-1}, \quad y > 0, \quad (3)$$

where $B(\cdot, \cdot)$ is the beta function.

Our goal is to derive a point predictor for $Y_{s:m}$ based on an observed sample, using the Bayesian approach under the ESEL function. For this purpose, let $\mathbf{X}_{(n)} = (X_1, \dots, X_n)$ denote the observed sample of size n from $Exp(\theta)$. The joint density function of $\mathbf{X}_{(n)}$ is given by

$$f(\mathbf{x}_{(n)}|\theta) = \theta^n e^{-\theta \sum_{i=1}^n x_i}, \quad (4)$$

where $\mathbf{x}_{(n)} = (x_1, \dots, x_n)$ is the observed set of $\mathbf{X}_{(n)}$. Here, we consider the conjugate prior distribution for θ , which is given as follows

$$\pi(\theta) = \frac{\beta^\alpha}{\Gamma(\alpha)} \theta^{\alpha-1} e^{-\beta\theta}, \quad \theta > 0, \quad \alpha, \beta > 0,$$

where $\Gamma(\cdot)$ is the gamma function, and α and β are hyperparameters that can be chosen based on the prior information. Thus, the posterior distribution is obtained to be

$$\pi(\theta|\mathbf{x}_{(n)}) = \frac{\left(\beta + \sum_{i=1}^n x_i\right)^{n+\alpha}}{\Gamma(n+\alpha)} \theta^{n+\alpha-1} e^{-\theta(\beta + \sum_{i=1}^n x_i)}. \quad (5)$$

Thus, $\theta|\mathbf{x}_{(n)}$ possesses a gamma distribution with parameters $n + \alpha$ and $\beta + \sum_{i=1}^n x_i$. So, from (3) and (5), the predictive density function for $Y_{s:m}$ is given by

$$\begin{aligned} f_{Y_{s:m}}(y|\mathbf{x}_{(n)}) &= \int_0^\infty f_{Y_{s:m}}(y;\theta) \pi(\theta|\mathbf{x}_{(n)}) d\theta \\ &= \frac{\left(\beta + \sum_{i=1}^n x_i\right)^{n+\alpha}}{B(s, m-s+1) \Gamma(n+\alpha)} \sum_{j=0}^{s-1} \binom{s-1}{j} (-1)^j \\ &\quad \times \int_0^\infty \theta^{n+\alpha} \exp\left\{-\left[(m-s+j+1)y + \beta + \sum_{i=1}^n x_i\right]\theta\right\} d\theta \\ &= \frac{\left(\beta + \sum_{i=1}^n x_i\right)^{n+\alpha} (n+\alpha)}{B(s, m-s+1)} \\ &\quad \times \sum_{j=0}^{s-1} \frac{\binom{s-1}{j} (-1)^j}{\left[(m-s+j+1)y + \beta + \sum_{i=1}^n x_i\right]^{n+\alpha+1}}. \end{aligned} \quad (6)$$

2.1 Point prediction

Let $Z_1, \dots, Z_m$ denote i.i.d. random variables from $Exp(\theta)$. Then, we have (Arnold et al., 2008)

$$Y_{s:m} \stackrel{d}{=} \sum_{l=1}^s \frac{Z_l}{m-l+1}, \quad (7)$$

where $\stackrel{d}{=}$ stands for identical in distribution. Therefore, from (7), we can write

$$E(Y_{s:m}) = \frac{1}{\theta} g(s, m), \quad \text{and} \quad Var(Y_{s:m}) = \frac{1}{\theta^2} h(s, m),$$

where

$$g(s, m) = \sum_{l=1}^s \frac{1}{m-l+1}, \quad \text{and} \quad h(s, m) = \sum_{l=1}^s \frac{1}{(m-l+1)^2}.$$

Note that the moment generating function (MGF) of Z_l , for $l = 1, \dots, m$ is given by

$$M_{Z_l}(t) = E(e^{tZ_l}) = \frac{\theta}{\theta - t}, \quad t < \theta.$$

Thus, from (7), the MGF of $Y_{s:m}$ is obtained to be

$$\begin{aligned}
 M_{Y_{s:m}}(t) &= E[e^{tY_{s:m}}] = E\left(\prod_{l=1}^s \exp\left\{t \frac{Z_l}{m-l+1}\right\}\right) \\
 &= \prod_{l=1}^s E\left(\exp\left\{t \frac{Z_l}{m-l+1}\right\}\right) \\
 &= \frac{\theta^s}{\prod_{l=1}^s \left(\theta - \frac{t}{m-l+1}\right)}, \quad t < \theta.
 \end{aligned} \tag{8}$$

The point predictor for $Y_{s:m}$ under the ESEL function is

$$\hat{Y}_{s:m} = -\log(E(e^{-Y_{s:m}} | \mathbf{X}_{(n)})).$$

Now, from (6) and (8), we get

$$\begin{aligned}
 E(e^{-Y_{s:m}} | \mathbf{x}_{(n)}) &= \int_0^\infty e^{-y} f_{Y_{s:m}}(y | \mathbf{x}_{(n)}) dy \\
 &= \int_0^\infty \int_0^\infty e^{-y} f_{Y_{s:m}}(y; \theta) dy \pi(\theta | \mathbf{x}_{(n)}) d\theta \\
 &= \int_0^\infty E(e^{-Y_{s:m}}) \pi(\theta | \mathbf{x}_{(n)}) d\theta \\
 &= \int_0^\infty M_{Y_{s:m}}(-1) \pi(\theta | \mathbf{x}_{(n)}) d\theta \\
 &= \int_0^\infty \frac{\theta^s}{\prod_{l=1}^s \left(\theta + \frac{1}{m-l+1}\right)} \times \frac{\left(\beta + \sum_{i=1}^n x_i\right)^{n+\alpha}}{\Gamma(n+\alpha)} \theta^{n+\alpha-1} e^{-\theta(\beta + \sum_{i=1}^n x_i)} d\theta \\
 &= E\left(\frac{\theta^s}{\prod_{l=1}^s \left(\theta + \frac{1}{m-l+1}\right)} \middle| \mathbf{x}_{(n)}\right).
 \end{aligned}$$

Consequently, the point prediction of $Y_{s:m}$ under the ESEL function is given by

$$\begin{aligned}
 \hat{Y}_{s:m} &= -\log E(e^{-Y_{s:m}} | \mathbf{x}_{(n)}) \\
 &= -\log \left[E\left(\frac{\theta^s}{\prod_{l=1}^s \left(\theta + \frac{1}{m-l+1}\right)} \middle| \mathbf{x}_{(n)}\right) \right].
 \end{aligned}$$

To approximate $\hat{Y}_{s:m}$, we generate a sample from $\theta | \mathbf{x}_{(n)}$, say $\{\theta_1, \dots, \theta_R\}$, where R is a large number. Then the approximate Bayesian prediction of $Y_{s:m}$ under the ESEL function is given by

$$\tilde{Y}_{s:m} = -\log \left(\frac{1}{R} \sum_{r=1}^R \frac{\theta_r^s}{\prod_{l=1}^s \left(\theta_r + \frac{1}{m-l+1}\right)} \right). \tag{9}$$

3 Results when N is a discrete random variable

In this section, we consider a scenario where the sample size itself is random. Let $\mathbf{X}_{(N)} = (X_1, \dots, X_N)$ denote the set of the sample observations, where N is a random variable and $\mathbf{X}_{(N)}$ consists of i.i.d. random variables coming from a one-parameter exponential distribution with parameter θ . We assume N is truncated at point $t_0^* = t_0 - 1$, namely N takes values in $\{t_0, t_0 + 1, t_0 + 2, \dots\}$. Furthermore, let N^* follow a power series distribution with the following probability mass function (pmf)

$$P_\lambda(N^* = n) = \frac{h(n)\lambda^n}{g(\lambda)}, \quad n = 0, 1, \dots,$$

where $g(\lambda) = \sum_{n=0}^{\infty} h(n)\lambda^n$. Then, the pmf of N can be stated as

$$P_\lambda(N = n) = \frac{h(n)\lambda^n}{P_\lambda(N^* \geq t_0)g(\lambda)}, \quad n = t_0, t_0 + 1, \dots$$

The class of power series distributions includes several well-known distributions, such as the Poisson, geometric, negative binomial, and logarithmic distributions, see, for example, Mahmoudi and Mahmoodian (2015). For further discussion on the properties, see Noack (1950) and Khatri (1959).

Let $\mathbf{x}_{(N)} = (x_1, \dots, x_N)$ denote the information sample with $x_1, \dots, x_n$ being the observed sample values. Then, from (4) and following Raghunandanan and Patil (1972), the joint pdf of $X_1, \dots, X_N$ is given by

$$f^R(\mathbf{x}_{(N)}|\theta) = \frac{1}{P(N^* \geq t_0)} \sum_{n=t_0}^{\infty} \theta^n e^{-\theta \sum_{i=1}^n x_i} P(N^* = n).$$

Thus, the posterior distribution of θ is given by

$$\pi^R(\theta|\mathbf{x}_{(N)}) = \frac{1}{P(N^* \geq t_0)} \sum_{n=t_0}^{\infty} \frac{\left(\beta + \sum_{i=1}^n x_i\right)^{n+\alpha}}{\Gamma(n+\alpha)} \theta^{n+\alpha-1} e^{-\theta(\beta + \sum_{i=1}^n x_i)} P(N^* = n).$$

The predictive density function for $Y_{s:m}$ is obtained to be

$$\begin{aligned} f_{Y_{s:m}}^R(y|\mathbf{x}_{(N)}) &= \frac{1}{B(s, m-s+1)P(N^* \geq t_0)} \sum_{n=t_0}^{\infty} (n+\alpha) \left(\beta + \sum_{i=1}^n x_i\right)^{n+\alpha} P(N^* = n) \\ &\quad \times \sum_{j=0}^{s-1} \frac{\binom{s-1}{j} (-1)^j}{\left[(m-s+j+1)y + \beta + \sum_{i=1}^n x_i\right]^{n+\alpha+1}}. \end{aligned}$$

Now, using a similar strategy stated in Subsection 2.1, we get

$$\begin{aligned} E(e^{-Y_{s:m}}|\mathbf{x}_{(N)}) &= \int_0^\infty E(e^{-Y_{s:m}}) \pi^R(\theta|\mathbf{x}_{(N)}) d\theta \\ &= \frac{1}{P(N^* \geq t_0)} \sum_{n=t_0}^{\infty} \int_0^\infty \frac{\theta^s}{\prod_{l=1}^s \left(\theta + \frac{1}{m-l+1}\right)} \end{aligned}$$

$$\begin{aligned}
& \times \frac{\left(\beta + \sum_{i=1}^n x_i\right)^{n+\alpha}}{\Gamma(n+\alpha)} \theta^{n+\alpha-1} e^{-\theta(\beta + \sum_{i=1}^n x_i)} d\theta P(N^* = n) \\
& = E_N \left\{ E \left(\frac{\theta^s}{\prod_{l=1}^s \left(\theta + \frac{1}{m-l+1} \right)} \middle| \mathbf{x}_{(N)}, N \right) \right\}.
\end{aligned}$$

Consequently, the point prediction of $Y_{s:m}$ under the ESEL function, when N is a random variable, is given by

$$\begin{aligned}
\hat{Y}_{s:m}^R &= -\log\{E(e^{-Y_{s:m}} | \mathbf{x}_{(N)})\} \\
&= -\log \left[E_N \left\{ E \left(\frac{\theta^s}{\prod_{l=1}^s \left(\theta + \frac{1}{m-l+1} \right)} \middle| \mathbf{x}_{(N)}, N \right) \right\} \right].
\end{aligned}$$

We may approximate $\hat{Y}_{s:m}^R$ using a similar approach stated in Subsection 2.1.

4 Simulation study

In this section, we carry out a simulation study to check the performance of the predictors discussed in the previous sections. We generate samples from an exponential distribution, $Exp(\theta)$, where two values are taken for θ as $\theta = 1$, and 3. We consider the following cases for the informative sample size:

- (i) Fixed sample size: $N = n$ is fixed, where n takes values 10, 20, 30, and 40.
- (ii) Random sample size (truncated Poisson): N follows a truncated Poisson distribution at $t_0^* = 1$, with parameter λ_0 selected such that $E_{\lambda_0}(N) = n$ for $n = 10, 20, 30$, and 40.
- (iii) Random sample size (truncated geometric): N follows a truncated geometric distribution at $t_0^* = 1$, with parameter p_0 chosen such that $E_{p_0}(N) = n$ for $n = 10, 20, 30$, and 40.

We examine three future sample sizes: $m = 15, 24$, and 35 with varying values of s . Two prior distributions are considered:

Prior 1 (non-informative): $a = b = 0$.

Prior 2 (informative): Hyperparameters a and b are selected such that the prior expected value of θ equals the true value of θ , and the prior variance is set to 2. Specifically, for $\theta = 1$, we have $a = b = 0.5$, and for $\theta = 3$, we have $a = 4.5$ and $b = 1.5$.

The number of the iterations of the simulation is $M = 10000$. Let $\hat{Y}_{s:m}$ be a predictor for $Y_{s:m}$. For the i -th replication, let $Y_{s:m}(i)$ be the generated s -th order statistic, and $\hat{Y}_{s:m}(i)$ be its corresponding predicted value. Then the estimated bias (bias for short), and the estimated risk (ER) of $\hat{Y}_{s:m}$ under the ESEL function are computed using the following relations:

$$bias(\hat{Y}_{s:m}) = \frac{1}{M} \sum_{i=1}^M (\hat{Y}_{s:m}(i) - Y_{s:m}(i)),$$

$$ER(\widehat{Y}_{s:m}) = \frac{1}{M} \sum_{i=1}^M (\exp\{-\widehat{Y}_{s:m}(i)\} - \exp\{-Y_{s:m}(i)\})^2,$$

respectively. We have computed the biases and ERs of the approximate point predictors, and the results are given in Tables 1–6. From these tables, we may draw the following conclusions:

Table 1: The ERs (first row) and the biases (second row) of the approximate Bayesian predictors under the ESEL function when $\theta = 1$, and $N = n$ is fixed.

		Prior 1				Prior 2			
m	$s \backslash n$	10	20	30	40	10	20	30	40
15	1	0.0223	0.0211	0.0207	0.0209	0.0221	0.0210	0.0206	0.0209
		0.0039	0.0011	0.0000	−0.0017	0.0036	0.0010	0.0000	−0.0017
	5	0.7063	0.6529	0.6332	0.6269	0.6986	0.6511	0.6323	0.6264
		0.0143	0.0022	−0.0042	−0.0090	0.0132	0.0020	−0.0044	−0.0091
	8	2.8338	2.5540	2.4440	2.3959	2.7882	2.5429	2.4390	2.3932
		0.0063	−0.0111	−0.0186	−0.0220	0.0051	−0.0114	−0.0188	−0.0221
	10	7.0346	6.1557	5.8227	5.6737	6.8813	6.1185	5.8062	5.6646
		−0.0184	−0.0304	−0.0419	−0.0472	−0.0193	−0.0307	−0.0421	−0.0473
	15	564.011	400.367	342.692	315.889	525.414	391.894	339.036	314.015
		−0.6365	−0.5773	−0.5908	−0.5917	−0.6299	−0.5756	−0.5901	−0.5913
24	1	0.0091	0.0087	0.0085	0.0085	0.0091	0.0087	0.0085	0.0085
		0.0034	0.0008	0.0004	−0.0007	0.0032	0.0008	0.0003	−0.0007
	5	0.2438	0.2255	0.2194	0.2164	0.2413	0.2250	0.2191	0.2163
		0.0160	0.0048	0.0020	0.0008	0.0151	0.0046	0.0019	0.0007
	10	1.4232	1.2768	1.2307	1.2106	1.4024	1.2721	1.2287	1.2093
		0.0217	0.0037	−0.0004	−0.0034	0.0202	0.0033	−0.0006	−0.0035
	12	2.5330	2.2395	2.1451	2.1083	2.4891	2.2296	2.1408	2.1057
		0.0198	−0.0007	−0.0027	−0.0083	0.0182	−0.0011	−0.0029	−0.0084
	15	6.0430	5.1721	4.9046	4.7955	5.9029	5.1409	4.8912	4.7874
		0.0083	−0.0118	−0.0132	−0.0202	0.0068	−0.0122	−0.0135	−0.0203
	20	39.3023	30.1697	27.6075	26.5623	37.5981	29.8146	27.4593	26.4731
		−0.0692	−0.0743	−0.0710	−0.0711	−0.0689	−0.0742	−0.0711	−0.0712
35	24	1965.23	1105.84	924.537	857.103	1772.65	1075.00	912.270	849.634
		−0.7064	−0.6386	−0.5880	−0.5986	−0.6981	−0.6361	−0.5870	−0.5980
	1	0.0043	0.0041	0.0040	0.0040	0.0043	0.0041	0.0040	0.0040
		0.0022	0.0013	0.0005	0.0001	0.0021	0.0013	0.0005	0.0001
	5	0.1076	0.1003	0.0973	0.0962	0.1067	0.1001	0.0972	0.0961
		0.0109	0.0049	0.0028	0.0015	0.0102	0.0047	0.0027	0.0014
	10	0.5274	0.4855	0.4699	0.4623	0.5215	0.4841	0.4692	0.4619
		0.0189	0.0081	0.0030	0.0016	0.0178	0.0078	0.0028	0.0016
	15	1.6130	1.4402	1.3826	1.3581	1.5878	1.4346	1.3801	1.3566
		0.0272	0.0073	0.0056	0.0028	0.0255	0.0069	0.0054	0.0027
	18	2.8873	2.5695	2.4552	2.3992	2.8367	2.5573	2.4498	2.3962
		0.0203	0.0080	−0.0032	−0.0040	0.0189	0.0076	−0.0034	−0.0041
	20	4.3804	3.7831	3.6002	3.5196	4.2867	3.7622	3.5912	3.5143
		0.0242	0.0040	−0.0007	−0.0031	0.0225	0.0035	−0.0009	−0.0033
	25	12.8344	10.8075	10.0445	9.7042	12.4584	10.7187	10.0055	9.6831
		−0.0061	−0.0144	−0.0225	−0.0250	−0.0067	−0.0146	−0.0227	−0.0251
	30	60.8660	46.9781	41.9444	39.6592	57.9331	46.3155	41.6575	39.5085
		−0.0832	−0.0756	−0.0749	−0.0759	−0.0817	−0.0753	−0.0748	−0.0759
	35	4603.63	2707.65	2120.37	1849.85	4105.37	2613.37	2081.07	1830.90
		−0.7966	−0.6892	−0.6471	−0.6397	−0.7851	−0.6860	−0.6457	−0.6389

Table 2: The ERs (first row) and the biases (second row) of the approximate Bayesian predictors under the ESEL function when $\theta = 3$, and $N = n$ is fixed.

m	$s \backslash n$	Prior 1				Prior 2			
		10	20	30	40	10	20	30	40
15	1	0.0026	0.0025	0.0025	0.0024	0.0025	0.0025	0.0025	0.0024
		0.0025	0.0009	0.0001	0.0006	0.0017	0.0006	0.0000	0.0005
	5	0.0791	0.0734	0.0719	0.0704	0.0749	0.0722	0.0714	0.0702
		0.0116	0.0047	0.0019	0.0014	0.0075	0.0036	0.0014	0.0011
	8	0.2787	0.2563	0.2487	0.2443	0.2614	0.2515	0.2465	0.2431
		0.0173	0.0058	0.0026	0.0006	0.0106	0.0038	0.0017	0.0001
	10	0.5807	0.5297	0.5125	0.5024	0.5400	0.5183	0.5072	0.4996
		0.0215	0.0064	0.0025	-0.0001	0.0131	0.0039	0.0014	-0.0006
	15	8.2157	7.0401	6.6607	6.4556	7.1447	6.7305	6.5164	6.3757
		-0.0419	-0.0591	-0.0585	-0.0682	-0.0499	-0.0618	-0.0598	-0.0687
24	1	0.0011	0.0010	0.0010	0.0010	0.0010	0.0010	0.0010	0.0010
		0.0013	0.0008	0.0003	0.0001	0.0008	0.0006	0.0002	0.0001
	5	0.0274	0.0254	0.0248	0.0247	0.0260	0.0251	0.0246	0.0246
		0.0078	0.0035	0.0023	0.0006	0.0052	0.0028	0.0020	0.0004
	10	0.1451	0.1334	0.1295	0.1283	0.1364	0.1310	0.1284	0.1277
		0.0156	0.0071	0.0041	0.0010	0.0102	0.0056	0.0034	0.0007
	12	0.2412	0.2204	0.2142	0.2115	0.2260	0.2162	0.2123	0.2105
		0.0186	0.0090	0.0046	0.0013	0.0120	0.0071	0.0037	0.0009
	15	0.4890	0.4446	0.4310	0.4236	0.4546	0.4350	0.4267	0.4213
		0.0244	0.0110	0.0047	0.0015	0.0159	0.0084	0.0036	0.0010
	20	1.7301	1.5329	1.4733	1.4396	1.5714	1.4882	1.4532	1.4287
		0.0279	0.0129	0.0022	-0.0014	0.0166	0.0094	0.0007	-0.0021
35	1	12.6596	10.5389	9.8689	9.5225	10.6924	9.9836	9.6173	9.3829
		-0.0425	-0.0513	-0.0683	-0.0704	-0.0498	-0.0540	-0.0694	-0.0707
	5	0.0005	0.0005	0.0005	0.0005	0.0005	0.0005	0.0005	0.0005
		0.0010	0.0004	0.0003	0.0002	0.0006	0.0004	0.0002	0.0002
	10	0.0123	0.0113	0.0110	0.0109	0.0116	0.0112	0.0110	0.0109
		0.0048	0.0024	0.0017	0.0009	0.0030	0.0019	0.0014	0.0008
	15	0.0572	0.0529	0.0515	0.0509	0.0540	0.0521	0.0511	0.0507
		0.0109	0.0049	0.0032	0.0017	0.0072	0.0039	0.0028	0.0015
	18	0.1599	0.1464	0.1419	0.1404	0.1502	0.1437	0.1407	0.1398
		0.0166	0.0082	0.0051	0.0022	0.0109	0.0066	0.0044	0.0018
	20	0.2670	0.2447	0.2382	0.2342	0.2497	0.2399	0.2360	0.2330
		0.0211	0.0096	0.0052	0.0027	0.0143	0.0076	0.0043	0.0022
35	25	0.3734	0.3388	0.3280	0.3239	0.3478	0.3317	0.3248	0.3222
		0.0226	0.0116	0.0065	0.0023	0.0147	0.0093	0.0055	0.0018
	30	0.8405	0.7576	0.7334	0.7180	0.7746	0.7391	0.7249	0.7133
		0.0296	0.0134	0.0065	0.0033	0.0197	0.0105	0.0051	0.0026
	35	2.1823	1.9306	1.8515	1.8008	1.9676	1.8696	1.8233	1.7854
		0.0335	0.0114	0.0035	0.0023	0.0220	0.0078	0.0018	0.0014
	35	17.6083	14.4325	13.4384	12.8776	14.5510	13.5591	13.0341	12.6541
		-0.0487	-0.0697	-0.0714	-0.0680	-0.0539	-0.0718	-0.0724	-0.0683

• The ERs exhibit a decreasing trend with respect to n in most cases, as expected, since increasing the number of observations should improve performance. Recall that n denotes the sample size when N is fixed, and the expected value of N when N is random (Clearly, we expect that increasing the expected value of N leads to larger values for the number of observations). Moreover, the absolute values of biases also decrease respect to n in most cases for $\theta = 1$ when N follows the truncated geometric distribution, and for $\theta = 3$ as well. The biases are negative when $s = m$.

Table 3: The ERs (first row) and the biases (second row) of the approximate Bayesian predictors under the ESEL function when $\theta = 1$, and N follows the truncated Poisson distribution at $t_0^* = 1$.

m	$s \backslash n$	Prior 1				Prior 2			
		10	20	30	40	10	20	30	40
15	1	0.0229	0.0212	0.0204	0.0203	0.0226	0.0212	0.0204	0.0203
		0.0045	0.0012	0.0009	0.0007	0.0040	0.0011	0.0008	0.0007
	5	0.7278	0.6565	0.6324	0.6269	0.7157	0.6543	0.6315	0.6264
		0.0183	0.0008	-0.0028	-0.0053	0.0167	0.0004	-0.0030	-0.0054
	8	2.9495	2.5626	2.4408	2.4090	2.8786	2.5496	2.4358	2.4062
		0.0074	-0.0114	-0.0159	-0.0201	0.0058	-0.0118	-0.0161	-0.0202
	10	7.3724	6.1839	5.8187	5.7037	7.1312	6.1399	5.8021	5.6943
		-0.0095	-0.0338	-0.0394	-0.0403	-0.0106	-0.0342	-0.0395	-0.0404
	15	636.061	404.007	339.888	319.355	571.712	393.827	336.283	317.407
		-0.6489	-0.6153	-0.5871	-0.5779	-0.6406	-0.6134	-0.5862	-0.5775
24	1	0.0095	0.0085	0.0085	0.0082	0.0094	0.0085	0.0085	0.0082
		0.0031	0.0016	0.0001	0.0008	0.0027	0.0016	0.0001	0.0008
	5	0.2511	0.2258	0.2202	0.2152	0.2471	0.2250	0.2199	0.2150
		0.0174	0.0065	0.0008	0.0014	0.0160	0.0062	0.0007	0.0013
	10	1.4711	1.2830	1.2327	1.2051	1.4374	1.2770	1.2305	1.2039
		0.0250	0.0070	-0.0026	-0.0028	0.0229	0.0065	-0.0028	-0.0029
	12	2.6307	2.2558	2.1510	2.0979	2.5596	2.2434	2.1464	2.0955
		0.0226	0.0018	-0.0074	-0.0072	0.0205	0.0012	-0.0076	-0.0073
	15	6.3327	5.2227	4.9145	4.7631	6.1032	5.1837	4.9001	4.7553
		0.0134	-0.0099	-0.0193	-0.0151	0.0116	-0.0105	-0.0195	-0.0152
35	1	0.0044	0.0042	0.0040	0.0040	0.0043	0.0042	0.0040	0.0040
		0.0035	0.0007	0.0005	0.0004	0.0032	0.0006	0.0005	0.0004
	5	0.1106	0.1007	0.0978	0.0960	0.1091	0.1004	0.0977	0.0960
		0.0132	0.0048	0.0021	0.0029	0.0122	0.0046	0.0020	0.0028
	10	0.5433	0.4871	0.4704	0.4630	0.5339	0.4855	0.4697	0.4626
		0.0246	0.0083	0.0030	0.0038	0.0229	0.0080	0.0028	0.0037
	15	1.6702	1.4492	1.3903	1.3575	1.6292	1.4421	1.3876	1.3561
		0.0325	0.0101	-0.0007	-0.0008	0.0303	0.0095	-0.0008	-0.0009
	18	3.0095	2.5814	2.4571	2.4094	2.9296	2.5669	2.4516	2.4064
		0.0280	0.0074	-0.0042	-0.0006	0.0261	0.0069	-0.0044	-0.0008
	20	4.5774	3.8142	3.6108	3.5082	4.4237	3.7882	3.6011	3.5030
		0.0314	0.0073	-0.0068	-0.0069	0.0293	0.0067	-0.0070	-0.0070
	25	13.6577	10.8591	10.0342	9.7725	13.0503	10.7531	9.9949	9.7507
		0.0035	-0.0124	-0.0207	-0.0210	0.0028	-0.0128	-0.0208	-0.0211
	30	66.8667	47.3147	41.7950	40.0258	61.9845	46.5196	41.5086	39.8694
		-0.0774	-0.0723	-0.0693	-0.0702	-0.0754	-0.0721	-0.0691	-0.0702
	35	5352.70	2752.22	2084.87	1869.51	4546.37	2637.57	2046.60	1849.82
		-0.8069	-0.7026	-0.6411	-0.6337	-0.7927	-0.6990	-0.6394	-0.6329

- As expected, Prior 2 (an informative prior) typically yields smaller ERs than Prior 1 (a non-informative prior), since we expect that an informative prior enhances performance. Furthermore, Prior 2 demonstrates superior performance in terms of bias in most cases, particularly when $\theta = 3$.
- For a fixed m , as s increases, the ERs increase. This is not far from our expectation, as we anticipate that with larger future observations, the error or loss would also in-

Table 4: The ERs (first row) and the biases (second row) of the approximate Bayesian predictors under the ESEL function when $\theta = 3$, and N follows the truncated Poisson distribution at $t_0^* = 1$.

m	$s \backslash n$	Prior 1				Prior 2			
		10	20	30	40	10	20	30	40
15	1	0.0027	0.0025	0.0025	0.0024	0.0025	0.0025	0.0024	0.0024
		0.0024	0.0012	0.0007	0.0006	0.0014	0.0010	0.0006	0.0005
	5	0.0795	0.0737	0.0718	0.0711	0.0743	0.0724	0.0712	0.0708
		0.0119	0.0050	0.0024	0.0010	0.0072	0.0037	0.0018	0.0007
	8	0.2797	0.2557	0.2499	0.2451	0.2589	0.2503	0.2474	0.2437
		0.0177	0.0083	0.0017	0.0016	0.0102	0.0062	0.0007	0.0010
	10	0.5837	0.5300	0.5147	0.5055	0.5357	0.5170	0.5090	0.5022
		0.0201	0.0092	0.0010	0.0001	0.0110	0.0064	-0.0003	-0.0006
	15	8.2120	7.0911	6.6815	6.5083	7.0023	6.7416	6.5254	6.4178
		-0.0496	-0.0548	-0.0680	-0.0660	-0.0546	-0.0577	-0.0694	-0.0667
24	1	0.0011	0.0010	0.0009	0.0010	0.0010	0.0010	0.0009	0.0010
		0.0017	0.0006	0.0007	0.0004	0.0010	0.0004	0.0006	0.0004
	5	0.0281	0.0254	0.0247	0.0246	0.0261	0.0250	0.0246	0.0245
		0.0084	0.0035	0.0020	0.0013	0.0050	0.0028	0.0017	0.0011
	10	0.1483	0.1333	0.1297	0.1285	0.1370	0.1307	0.1286	0.1278
		0.0173	0.0064	0.0031	0.0020	0.0103	0.0048	0.0024	0.0016
	12	0.2459	0.2208	0.2145	0.2124	0.2262	0.2162	0.2126	0.2112
		0.0214	0.0072	0.0033	0.0019	0.0131	0.0053	0.0025	0.0013
	15	0.5011	0.4453	0.4306	0.4258	0.4569	0.4348	0.4262	0.4231
		0.0259	0.0083	0.0038	0.0020	0.0155	0.0058	0.0027	0.0013
	20	1.7717	1.5343	1.4710	1.4502	1.5725	1.4857	1.4503	1.4375
		0.0318	0.0073	0.0009	-0.0017	0.0186	0.0040	-0.0005	-0.0027
35	1	13.1194	10.5232	9.8319	9.6013	10.6848	9.4148	9.5717	9.4430
		-0.0368	-0.0639	-0.0668	-0.0705	-0.0437	-0.0657	-0.0675	-0.0712
	5	0.0005	0.0005	0.0005	0.0005	0.0005	0.0005	0.0005	0.0005
		0.0009	0.0005	0.0002	0.0001	0.0005	0.0004	0.0002	0.0001
	10	0.0123	0.0114	0.0111	0.0110	0.0115	0.0112	0.0110	0.0109
		0.0053	0.0024	0.0014	0.0008	0.0032	0.0018	0.0012	0.0007
	15	0.0576	0.0532	0.0517	0.0510	0.0537	0.0522	0.0513	0.0508
		0.0110	0.0052	0.0030	0.0021	0.0067	0.0040	0.0025	0.0018
	18	0.1639	0.1462	0.1422	0.1412	0.1511	0.1433	0.1410	0.1405
		0.0176	0.0076	0.0042	0.0023	0.0103	0.0059	0.0035	0.0018
	20	0.2688	0.2466	0.2390	0.2352	0.2481	0.2411	0.2366	0.2339
		0.0205	0.0092	0.0046	0.0032	0.0128	0.0070	0.0036	0.0026
35	25	0.3817	0.3385	0.3283	0.3248	0.3489	0.3309	0.3251	0.3228
		0.0247	0.0102	0.0052	0.0035	0.0149	0.0078	0.0043	0.0028
	30	0.8462	0.7629	0.7348	0.7220	0.7683	0.7418	0.7254	0.7167
		0.0271	0.0133	0.0061	0.0039	0.0167	0.0100	0.0046	0.0031
	35	2.1954	1.9433	1.8543	1.8150	1.9464	1.8741	1.8236	1.7974
		0.0268	0.0121	0.0035	0.0020	0.0156	0.0081	0.0016	0.0010
	35	17.5953	14.5588	13.4498	13.0262	14.1736	13.5749	13.0104	12.7718
		-0.0742	-0.0576	-0.0613	-0.0675	-0.0745	-0.0597	-0.0624	-0.0680

crease. However, this result does not consistently extend to bias.

- The approximate predictors with fixed sizes achieve the best performance in terms of ER in most cases. Moreover, the best performance in terms of bias generally occurs for either the approximate predictors with fixed sizes or those with truncated Poisson distributed sizes. In contrast, the approximate predictors with truncated geometrically distributed sample sizes exhibit the poorest performance in terms of both ER and bias.

Table 5: The ERs (first row) and the biases (second row) of the approximate Bayesian predictors under the ESEL function when $\theta = 1$, and N follows the truncated geometric distribution at $t_0^* = 1$.

m	$s \backslash n$	Prior 1				Prior 2			
		10	20	30	40	10	20	30	40
15	1	0.0254	0.0234	0.0226	0.0221	0.0243	0.0228	0.0222	0.0218
		0.0107	0.0064	0.0034	0.0032	0.0090	0.0054	0.0028	0.0026
	5	0.8449	0.7543	0.7089	0.6996	0.7966	0.7288	0.6919	0.6845
		0.0343	0.0178	0.0094	0.0063	0.0301	0.0151	0.0077	0.0049
	8	3.5408	3.0580	2.8249	2.7841	3.2683	2.9159	2.7315	2.6951
		0.0254	0.0094	-0.0019	0.0015	0.0223	0.0070	-0.0032	0.0002
	10	9.2309	7.6902	6.9808	6.9245	8.3135	7.2223	6.6760	6.6133
		-0.0013	-0.0156	-0.0248	-0.0196	-0.0021	-0.0169	-0.0254	-0.0204
	15	1062.609	665.534	535.685	609.055	808.074	556.900	472.060	506.672
		-0.7161	-0.6580	-0.6322	-0.6213	-0.6926	-0.6469	-0.6239	-0.6152
24	1	0.0103	0.0095	0.0089	0.0089	0.0098	0.0093	0.0088	0.0088
		0.0072	0.0042	0.0037	0.0029	0.0059	0.0036	0.0033	0.0025
	5	0.2823	0.2551	0.2405	0.2386	0.2675	0.2480	0.2357	0.2338
		0.0276	0.0172	0.0123	0.0079	0.0238	0.0151	0.0110	0.0067
	10	1.6836	1.4867	1.3866	1.3683	1.5697	1.4332	1.3498	1.3287
		0.0367	0.0206	0.0149	0.0087	0.0328	0.0183	0.0135	0.0075
	12	3.0543	2.6518	2.4568	2.4268	2.8162	2.5421	2.3810	2.3402
		0.0320	0.0171	0.0125	0.0066	0.0288	0.0151	0.0113	0.0055
	15	7.6054	6.3586	5.8172	5.8071	6.8378	6.0197	5.5802	5.5016
		0.0149	0.0034	-0.0004	-0.0075	0.0139	0.0024	-0.0009	-0.0080
	20	58.4913	42.1403	37.0691	40.3946	48.0637	38.2021	34.2446	35.2391
		-0.0997	-0.0762	-0.0741	-0.0829	-0.0926	-0.0729	-0.0715	-0.0809
	24	4656.21	2121.87	1731.98	2985.15	3140.42	1743.49	1458.78	2117.03
		-0.8318	-0.6955	-0.6822	-0.6623	-0.8007	-0.6792	-0.6701	-0.6524
35	1	0.0049	0.0047	0.0043	0.0045	0.0047	0.0045	0.0043	0.0044
		0.0066	0.0034	0.0028	0.0013	0.0055	0.0028	0.0024	0.0010
	5	0.1272	0.1139	0.1087	0.1069	0.1203	0.1103	0.1063	0.1048
		0.0259	0.0169	0.0096	0.0086	0.0224	0.0149	0.0083	0.0075
	10	0.6361	0.5627	0.5303	0.5223	0.5975	0.5422	0.5168	0.5102
		0.0416	0.0274	0.0153	0.0134	0.0369	0.0245	0.0135	0.0118
	15	1.9268	1.6867	1.5735	1.5512	1.7879	1.6220	1.5289	1.5018
		0.0432	0.0277	0.0181	0.0146	0.0393	0.0254	0.0167	0.0133
	18	3.6740	3.1245	2.8778	2.8436	3.3630	2.9637	2.7724	2.7410
		0.0471	0.0319	0.0142	0.0143	0.0435	0.0292	0.0127	0.0128
	20	5.4594	4.5988	4.2342	4.2067	4.9432	4.3689	4.0736	4.0060
		0.0331	0.0232	0.0116	0.0084	0.0310	0.0217	0.0107	0.0076
	25	18.1454	14.3189	12.7094	12.7688	15.7952	13.1604	11.9699	11.9364
		0.0092	0.0068	-0.0151	-0.0073	0.0102	0.0062	-0.0152	-0.0077
	30	100.7667	70.5845	59.4174	62.8730	81.5130	61.7865	54.0542	55.4958
		-0.0966	-0.0665	-0.0835	-0.0685	-0.0878	-0.0631	-0.0808	-0.0668
	35	11098.4	5386.91	4033.66	5468.67	7598.65	4164.74	3356.85	4030.67
		-0.9081	-0.7777	-0.7441	-0.7047	-0.8719	-0.7596	-0.7308	-0.6948

5 Real data example

For illustrative purposes, we consider real-world data, focusing on the lifetimes of steel specimens (in units of 1000 cycles) subjected to different stress amplitudes, see Maennig (1967); Kimber (1990); Crowder (2000); Lawless (2003). The specimens were tested under 14 distinct stress levels, ranging from 32.0 to 38.5 in increments of 0.5

Table 6: The ERs (first row) and the biases (second row) of the approximate Bayesian predictors under the ESEL function when $\theta = 3$, and N follows the truncated geometric distribution at $t_0^* = 1$.

m	$s \backslash n$	Prior 1				Prior 2			
		10	20	30	40	10	20	30	40
15	1	0.0031	0.0028	0.0027	0.0026	0.0026	0.0026	0.0025	0.0025
		0.0051	0.0025	0.0021	0.0016	0.0020	0.0009	0.0009	0.0007
	5	0.0933	0.0827	0.0802	0.0770	0.0766	0.0744	0.0734	0.0722
		0.0236	0.0128	0.0107	0.0092	0.0099	0.0056	0.0050	0.0050
	8	0.3297	0.2876	0.2800	0.2693	0.2670	0.2567	0.2543	0.2512
		0.0357	0.0207	0.0162	0.0129	0.0154	0.0100	0.0075	0.0064
	10	0.6936	0.5986	0.5844	0.5587	0.5530	0.5296	0.5262	0.5179
		0.0417	0.0248	0.0173	0.0150	0.0179	0.0122	0.0069	0.0072
	15	10.6917	8.4689	8.2536	7.6475	7.2515	6.9079	6.8198	6.7087
		-0.0261	-0.0469	-0.0509	-0.0539	-0.0401	-0.0534	-0.0588	-0.0601
24	1	0.0013	0.0011	0.0010	0.0010	0.0010	0.0010	0.0010	0.0010
		0.0031	0.0021	0.0017	0.0010	0.0010	0.0010	0.0009	0.0004
	5	0.0326	0.0289	0.0272	0.0268	0.0264	0.0258	0.0252	0.0252
		0.0158	0.0095	0.0069	0.0051	0.0063	0.0044	0.0035	0.0024
	10	0.1723	0.1519	0.1430	0.1401	0.1381	0.1345	0.1318	0.1307
		0.0308	0.0187	0.0134	0.0105	0.0135	0.0093	0.0071	0.0054
	12	0.2879	0.2523	0.2371	0.2321	0.2292	0.2224	0.2179	0.2159
		0.0356	0.0220	0.0157	0.0123	0.0155	0.0111	0.0084	0.0064
	15	0.5891	0.5115	0.4798	0.4693	0.4614	0.4467	0.4383	0.4340
		0.0435	0.0276	0.0185	0.0143	0.0199	0.0146	0.0099	0.0073
	20	2.1550	1.8130	1.6784	1.6308	1.5904	1.5342	1.5027	1.4779
		0.0485	0.0301	0.0165	0.0152	0.0225	0.0153	0.0069	0.0073
	24	18.0989	13.5386	11.8956	11.6180	10.7653	10.2779	9.9670	9.7913
		-0.0470	-0.0505	-0.0561	-0.0631	-0.0538	-0.0552	-0.0582	-0.0650
35	1	0.0006	0.0005	0.0005	0.0005	0.0005	0.0005	0.0005	0.0005
		0.0024	0.0014	0.0009	0.0010	0.0009	0.0007	0.0003	0.0005
	5	0.0144	0.0127	0.0125	0.0118	0.0118	0.0114	0.0114	0.0111
		0.0117	0.0069	0.0050	0.0049	0.0048	0.0033	0.0022	0.0028
	10	0.0678	0.0597	0.0581	0.0558	0.0551	0.0534	0.0529	0.0521
		0.0228	0.0133	0.0105	0.0089	0.0100	0.0066	0.0051	0.0049
	15	0.1910	0.1688	0.1576	0.1544	0.1525	0.1489	0.1449	0.1438
		0.0319	0.0184	0.0141	0.0111	0.0138	0.0085	0.0076	0.0058
	18	0.3196	0.2781	0.2693	0.2594	0.2563	0.2467	0.2433	0.2410
		0.0388	0.0232	0.0194	0.0152	0.0178	0.0121	0.0104	0.0085
	20	0.4489	0.3927	0.3659	0.3582	0.3529	0.3436	0.3347	0.3317
		0.0413	0.0242	0.0184	0.0143	0.0187	0.0117	0.0101	0.0075
	25	1.0198	0.8733	0.8454	0.8099	0.7931	0.7624	0.7512	0.7441
		0.0518	0.0296	0.0254	0.0185	0.0255	0.0157	0.0138	0.0098
	30	2.7224	2.2645	2.1938	2.0770	2.0103	1.9255	1.8973	1.8747
		0.0569	0.0301	0.0256	0.0175	0.0300	0.0159	0.0131	0.0080
	35	25.1253	18.3013	17.9623	16.0743	14.7131	13.9189	13.6906	13.4482
		-0.0393	-0.0630	-0.0473	-0.0591	-0.0434	-0.0637	-0.0516	-0.0626

(i.e. 32.0, 32.5, 33.0, ..., 38.0, 38.5). Prior studies have investigated different subsets of the stress levels tested. Taniş et al. (2023) focused on the stress levels 32.0 and 33.0, whereas Etemad Golestani et al. (2025) analyzed the data for 33.0 and 32.5 stress levels. In this work, we specifically analyze the data for stress level 32.0 and stress level 32.5, which are divided by 1000.

We used the Kolmogorov-Smirnov (K-S) test to check if the data come from the

exponential distribution with $\theta = 0.75$. The K-S test statistics are computed as $D_1 = 0.14316$ for the stress level 32.0, and $D_2 = 0.13671$ for the stress level 32.5. Moreover, the corresponding computed p -values are given by $p_1 = 0.6571$ for the stress level 32.0, and $p_2 = 0.8011$ for the stress level 32.5. We treat the first data (the stress level 32.0) as the information sample, denoted by $\mathbf{x}_{(n)} = (x_1, \dots, x_n)$, and the second data (the stress level 32.5) as the future sample. Assuming that there is no prior information, let us take $a = b = 0$. We define the deviation as $|\tilde{Y}_{s:m} - Y_{s:m}|$. From (2), we define the loss as $(e^{-\tilde{Y}_{s:m}} - e^{-Y_{s:m}})^2$. Table 7 presents several order statistics from the future sample along with their corresponding approximate Bayes point predictions under the ESEL function, computed using (9), as well as the deviations and the losses. As shown in Table 7, the predicted values approximately deviate from the actual values in an increasing manner as s grows. However, the loss values seem to remain stable.

Table 7: Several order statistics, and the approximate Bayes point predictions under the ESEL function, the deviations and the losses for data of stress level 32.5.

s	1	3	5	8	10	12	15	18	20
$Y_{s:m}$	0.196	0.250	0.308	0.475	0.669	0.879	1.338	2.211	4.257
$\tilde{Y}_{s:m}$	0.070	0.222	0.390	0.682	0.914	1.189	1.734	2.628	3.955
Deviation	0.126	0.028	0.082	0.207	0.245	0.310	0.396	0.417	0.302
Loss	0.012	0.000	0.003	0.014	0.012	0.012	0.007	0.001	0.000

Discussion and conclusions

In this paper, we addressed the problem of Bayesian two-sample prediction for the exponential distribution using a recently introduced asymmetric loss function. We examine cases where the informative sample size is either fixed or follows a random distribution (e.g., truncated Poisson or truncated geometric). Due to the fact that the Bayesian predictors do not seem to possess a closed form, we adopt a Monte Carlo approximation approach. Through a simulation study, we assess the performance of these predictors and demonstrate that fixed sample sizes may yield better predictive results compared to truncated Poisson and and truncated geometrically distributed sample sizes. A real data application is provided to illustrate the methodology. Note that other loss functions could be considered for predicting future observations. All computations in this paper were performed using the statistical software R (R Core Team, 2025) and the package extraDistr (Wolodzko, 2023) included therein.

Acknowledgement

The authors would like to express their gratitude to the reviewers for their valuable comments.

References

- Ahmadi, J., Basiri, E. and MirMostafaei, S.M.T.K. (2016). Optimal random sample size based on Bayesian prediction of exponential lifetime and application to real data. *Journal of the Korean Statistical Society*, **45**(2):221–237.
- AL-Hussaini, E.K. and Al-Awadhi, F. (2010). Bayes two-sample prediction of generalized order statistics with fixed and random sample size. *Journal of Statistical Computation and Simulation*, **80**(1):13–28.
- Arnold, B.C., Balakrishnan, N. and Nagaraja, H.N. (2008). *A First Course in Order Statistics*. Philadelphia: Society for Industrial and Applied Mathematics.
- Asgharzadeh, A. and Valiollahi, R. (2012). Prediction of times to failure of censored units in hybrid censored samples from exponential distribution. *Journal of Statistical Research of Iran*, **9**(1):11–30.
- Ashour, S.K. and El-Wekeel, M.A.M.H. (1993). Bayesian prediction of the j th order statistic with Burr distribution and a random sample size. *Microelectronics Reliability*, **33**(8):1179–1188.
- Barakat, H., Khaled, O. and Ghonem, H. (2021). Predicting future order statistics with random sample size. *AIMS Mathematics*, **6**(5):5133–5147.
- Barakat, H.M., Nigm, E.M., El-Adll, M.E. and Yusuf, M. (2018). Prediction of future generalized order statistics based on exponential distribution with random sample size. *Statistical Papers*, **59**(2):605–631.
- Basiri, E. and Ahmadi, J. (2015). Prediction intervals for generalized-order statistics with random sample size. *Journal of Statistical Computation and Simulation*, **85**(9):1725–1741.
- Basiri, E. and Pakzad, A. (2018). Optimal non-parametric prediction intervals for order statistics with random sample size. *Interdisciplinary Journal of Management Studies (Formerly known as Iranian Journal of Management Studies)*, **11**(2):307–322.
- Buhrman, J.M. (1973). On order statistics when the sample size has a binomial distribution. *Statistica Neerlandica*, **27**(3):125–126.
- Consul, P.C. (1984). On the distributions of order statistics for a random sample size. *Statistica Neerlandica*, **38**(4):249–256.
- Crowder, M. (2000). Tests for a family of survival models based on extremes. In *Eds. N. Limnios and M. Nikulin, Recent Advances in Reliability Theory: Methodology, Practice, and Inference*, Boston, Birkhäuser, 307–321.
- Etemad Golestani, B., Ormoz, E. and MirMostafaei, S.M.T.K. (2025). On the estimation of the stress-strength parameter for the inverse Lindley distribution based on lower record values. *Hacettepe Journal of Mathematics & Statistics*, **54**(1):291–318.

- Fallah, A. and Asgharzadeh, A. (2021). Pivotal and Bayesian inference in exponential coherent systems under progressive censoring. *Journal of Advanced Mathematical Modeling*, **11**(3):496–514.
- Gupta D. and Gupta R.C. (1984). On the distribution of order statistics for a random sample size. *Statistica Neerlandica*, **38**(1):13–19.
- Kaminsky, K.S. and Nelson, P.I. (1975). Best linear unbiased prediction of order statistics in location and scale families. *Journal of the American Statistical Association*, **70**(349):145–150.
- Khatri, C.G. (1959). On certain properties of power-series distributions. *Biometrika*, **46**(3/4):486–490.
- Kimber, A.C. (1990). Exploratory data analysis for possibly censored data from skewed distributions. *Journal of the Royal Statistical Society Series C (Applied Statistics)*, **39**(1):21–30.
- Kumar, D., Kumar, P., Kumar, P., Singh, U. and Chaurasia, P.K. (2020). A new asymmetric loss function for estimation of any parameter. *International Journal of Agricultural & Statistical Sciences*, **16**(2):489–501.
- Kumar, P., Kumar, P., Kumar, D., Singh, S.K. and Singh, U. (2023). Bayesian estimation of parameter for different loss functions using progressive Type-II censored data. *Electronic Journal of Applied Statistical Analysis*, **16**(3):519–540.
- Lawless, J.F. (1971). A prediction problem concerning samples from the exponential distribution, with application in life testing. *Technometrics*, **13**(4):725–730.
- Lawless, J.F. (1977). Prediction intervals for the two parameter exponential distribution. *Technometrics*, **19**(4):469–472.
- Lawless, J.F. (2003). *Statistical Models and Methods for Lifetime Data*. Second Edition, Hoboken: John Wiley & Sons.
- Lingappaiah, G.S. (1978). Bayesian approach to the prediction problem in the exponential population. *IEEE Transactions on Reliability*, **R-27**(3):222–225.
- Lingappaiah, G.S. (1979). Bayesian approach to the prediction problem in complete and censored samples from the gamma and exponential populations. *Communications in Statistics-Theory and Methods*, **8**(14):1403–1424.
- Lingappaiah, G.S. (1990). Bayesian inference in life tests based on exponential model with outliers when sample size is a random variable. *Trabajos de Estadística*, **5**(1):27–37.
- Maennig, W.W. (1967). Bemerken zur Beurteilung des dauerschwing Festigkeitsverhaltens von Stahl und einige Untersuchungen zer Bestimmig des Dauerfestigkeitsbereichs. *Materialpruf*, **12**:124–131.
- Mahmoudi, E. and Mahmoodian, H. (2015). Normal power series class of distributions: Model, properties and applications. *arXiv preprint arXiv:1510.07180*.

- MirMostafaei, S.M.T.K. and Ahmadi, J. (2011). Point prediction of future order statistics from an exponential distribution. *Statistics & Probability Letters*, **81**(3):360–370.
- Nigm, A.M. and Abd Al-Wahab, N.Y. (1996). Bayesian prediction with a random sample size for the Burr lifetime distribution. *Communications in Statistics-Theory and Methods*, **25**(6):1289–1303.
- Noack, A. (1950). A class of random variable with discrete distribution. *The Annals of Mathematical Statistics*, **21**(1):127–132.
- R Core Team (2025). R: A language and environment for statistical computing. Vienna, Austria: R Foundation for Statistical Computing. <https://www.R-project.org>
- Raghunandanan, K. and Patil, S.A. (1972). On order statistics for random sample size. *Statistica Neerlandica*, **26**(4):121–126.
- Rohatgi, V.K. (1987). Distribution of order statistics with random sample size. *Communications in Statistics-Theory and Methods*, **16**(12):3739–3743.
- Shafay, A.R., Balakrishnan, N. and Ahmadi, J. (2017). Bayesian prediction of order statistics with fixed and random sample sizes based on k -record values from Pareto distribution. *Communications in Statistics-Theory and Methods*, **46**(2):721–735.
- Soliman, A.A. (2000). Bayes prediction in a Pareto lifetime model with random sample size. *Journal of the Royal Statistical Society Series D (The Statistician)*, **49**(1):51–62.
- Srivastava, R.C. (1973). Estimation of probability density function based on random number of observations with applications. *International Statistical Review / Revue Internationale de Statistique*, **41**(1):77–86.
- Taniş, C., Saraçoğlu, B., Asgharzadeh, A. and Abdi, M. (2023). Estimation of $Pr(X < Y)$ for exponential power records. *Hacettepe Journal of Mathematics & Statistics*, **52**(2):499–511.
- Wolodzko, T. (2023). extraDistr: Additional univariate and multivariate distributions, R package version 1.10.0, <https://CRAN.R-project.org/package=extraDistr>.